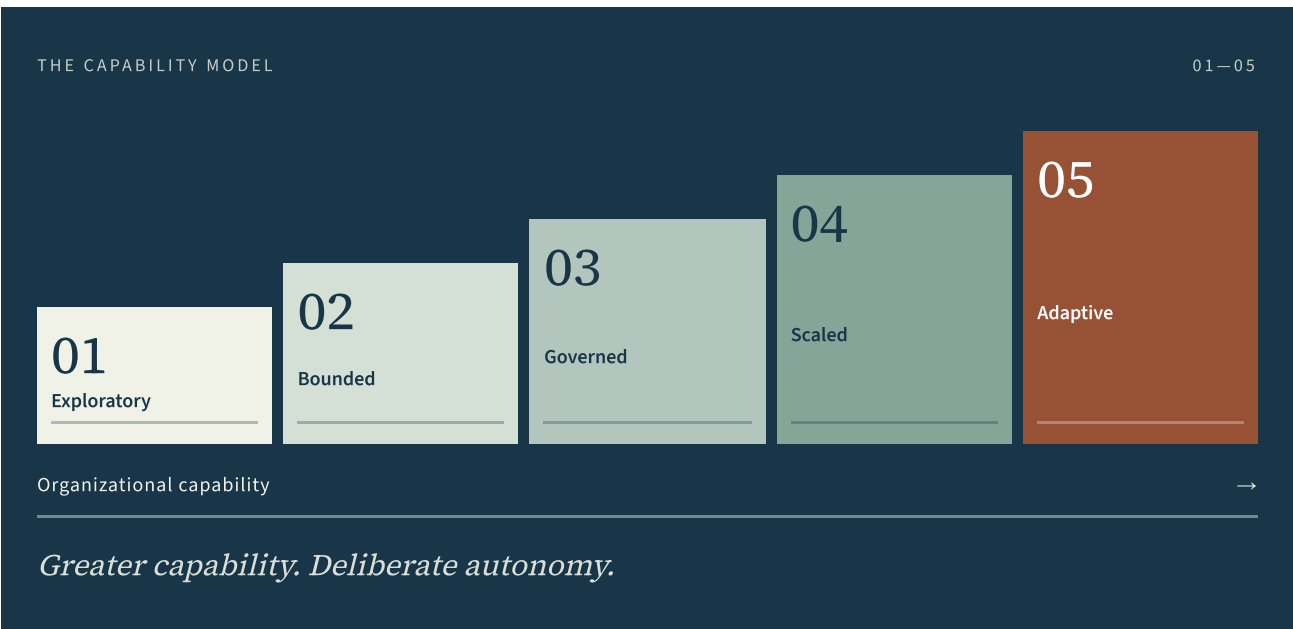


Enterprise Agentic Transformation Maturity Model

An evidence-led model for assessing whether an organization can delegate work to agents responsibly, operate it reliably, and realize measurable value.

By [Balaji Krishnan](#) Published September 6, 2026 Updated 2026-09-06 7 min read



Abstract

An organization can operate a sophisticated agent while remaining poorly prepared to govern the work it performs. This model assesses the organizational capabilities around agentic systems: ownership, data, architecture, controls, operations and value. Its unit of assessment is a defined business workflow in an organizational context. It is not a ranking of models, vendors or individual employees.

The model introduces five maturity levels across eight dimensions. Progress depends on observable practices and retained evidence. More autonomy does not imply greater maturity: a well-governed workflow with mandatory human approval may be more mature than an autonomous deployment whose owner cannot reconstruct a failed action.

Definition and scope

Enterprise agentic transformation is the deliberate redesign of work, decision rights, technology and controls so that an organization can use systems that select and execute actions toward an objective. For assessment purposes, an agent is a system in which a model can select subsequent actions or tool use within delegated authority. A fixed automation may be part of the same workflow, but it should not be classified as agentic solely because it calls a language model.

Use the model to assess a service, workflow or bounded portfolio. Name the scope, environment, assessment date and accountable executive. Do not extrapolate one successful pilot into a score for an entire enterprise. Record where evidence is unavailable and which teams or systems are outside the assessment boundary.

Methodology and limitations

This is an original conceptual assessment framework by Balaji Krishnan. It is informed by NIST's voluntary AI Risk Management Framework, OWASP's analysis of excessive agency and the workflow/agent distinction described by Anthropic. The dimensions, level definitions and scoring rules below are this publication's synthesis. They are not NIST or OWASP scoring instruments.

Version 1.0 has not been statistically validated against organizational outcomes. No benchmark population or predictive claim is implied. The levels are ordinal descriptions, not equal intervals. A mean score may help organize a workshop, but cannot establish readiness, compliance or safety.

The design separates business direction from operating responsibility, architecture from data readiness, and security from broader governance. This allows different accountable owners to contribute evidence without treating the platform team as the owner of every transformation problem.

The eight dimensions

Dimension	Assessment question	Minimum useful evidence
Strategy and scope	Is the delegated work bounded and connected to a business need?	Workflow charter, exclusions, accountable sponsor, baseline
People and operating model	Who owns the agent's work and handles exceptions?	Decision-rights map, staffed escalation rota, training records

Dimension	Assessment question	Minimum useful evidence
Data readiness	Is the necessary information authorized, fit for purpose and current?	Data inventory, access tests, freshness criteria, retention policy
Architecture and integration	Can tool calls and system effects be controlled independently of the model?	System boundary diagram, permission design, sandbox tests
Governance and assurance	Can reviewers understand why the deployment is allowed?	Agent register, approval record, evaluation plan, change log
Security and risk	Are abuse paths and consequential actions constrained?	Threat model, adversarial tests, least-privilege access, risk acceptance
Operations and reliability	Can the team detect, contain and recover from failure?	Service objectives, traces, incident exercises, rollback procedure
Value realization	Can the organization distinguish useful outcomes from activity?	Baseline comparison, cost ledger, quality measures, benefits owner

Five maturity levels

- 1. Exploratory.** Work is experimental. Boundaries and ownership may be provisional. Evidence consists mainly of hypotheses and demonstrations. Keep activity in a sandbox without consequential production authority.
- 2. Bounded.** A workflow has an owner, explicit exclusions, controlled data access and a test plan. The team can describe the intended failure response. Limited pilots may be considered after local review; the label alone does not authorize deployment.
- 3. Governed.** Approval, evaluation and operational practices are repeatable for the assessed workflow. There is a record of tool permissions, change review, incident handling and baseline outcomes. Reviewers can reconstruct sampled decisions and actions.
- 4. Scaled.** Shared controls support multiple instances or workflows without depending on exceptional effort by a few individuals. Capacity, costs, permission drift and cross-system failures are monitored. Exceptions have owners and closure dates.

5. Adaptive. Evidence drives controlled changes to the operating model. The organization can retire ineffective agents, revise delegation and compare alternatives. Evaluation and incident findings change policy through an accountable process. This level does not mean unbounded autonomy.

Assessment guidance

Convene the workflow owner, platform lead, risk or assurance representative, security lead and frontline operator. Score each dimension independently before discussing differences. For each proposed score, name an artifact, its owner and a recent example of the practice in use.

Award the highest level whose practices are evidenced consistently, including the expectations of earlier levels. A policy without operating evidence should not be scored as a repeatable practice. Treat unknown evidence as a gap; do not substitute optimism for a score.

Use the accompanying readiness assessment to record a profile. It reports the lowest dimension as the limiting level and a separate descriptive average. The average is never a production approval. Before any consequential deployment, require named ownership, enforceable access boundaries, tested escalation and a recovery plan regardless of the score.

Evidence indicators by level

Level	Evidence that supports progression	Evidence that challenges the score
Exploratory	Clear hypothesis and isolated experiment	Undocumented production access
Bounded	Signed scope, tested permissions, named exception owner	Exclusions exist only in a prompt
Governed	Repeated evaluations, sampled traces, exercised response	Approval record without incident capability
Scaled	Shared control coverage, drift detection, cost attribution	Manual exceptions multiplying faster than review capacity
Adaptive	Documented changes based on outcomes, retirement decisions	Activity metrics replacing quality or business results

A practical transformation roadmap

First 30 days: establish the boundary. Select a workflow, document its baseline and identify consequential actions. Map information and permissions. Assign the business owner and escalation responsibility. Record the initial evidence profile.

Next 30 days: prove the controls. Exercise the workflow using representative cases and deliberate failure scenarios. Test denied access, unavailable tools, duplicate requests and interrupted execution. Agree what evidence would justify a pilot and what would stop it.

Next 30 days: evaluate a bounded deployment. Compare outcomes with the baseline and account for reviewer workload, failures and operating cost. Make a documented decision to expand, revise or stop. These time windows are workshop planning suggestions, not evidence-based delivery promises.

Practical applications

An architecture board can use the dimensions to identify missing control owners before reviewing a design. A CIO can use workflow profiles to prioritize shared capabilities such as identity or incident response. A transformation lead can turn low-scoring dimensions into specific investment decisions. A risk team can identify the evidence it needs without prescribing a particular model vendor.

Do not use the model to rank employees, assert regulatory conformity or compare organizations with different assessment scopes. When presenting a portfolio, retain each workflow's boundary and evidence gaps alongside any aggregate view.

Version history

1.0 — 6 September 2026. Initial conceptual release: eight dimensions, five maturity levels, evidence-led assessment and a bounded implementation roadmap. Future minor releases will clarify guidance; changes to dimensions or scoring semantics require a new major version. Earlier versions remain available at their versioned URLs.

Citation and reuse

Krishnan, B. (2026). *Enterprise Agentic Transformation Maturity Model* (Version 1.0).

AgenticTransformation.ai. Cite the versioned URL and your access date. A DOI has not been assigned; the stable versioned URL is the current identifier.

You may reference and link to the model with attribution. Public redistribution or adaptation of the full work requires permission unless a separate license is issued. Do not imply endorsement or validation by this publication.

Adoption evidence

No independently verified adoption records have been published for this release. The adoption index will link to permission-cleared evidence as it becomes available. Downloads and visits will not be presented as proof of organizational use.

Stable identifier: <https://agentictransformation.ai/frameworks/enterprise-agentic-transformation-maturity-model/v1.0/>

References

1. [NIST: AI Risk Management Framework ↗](#)

Published 2023-01-26. Reviewed 2026-09-06. Voluntary framework for incorporating trustworthiness into AI design, development, use and evaluation. Framework publication date; landing page is maintained.

2. [Anthropic: Building effective agents ↗](#)

Published 2024-12-19. Reviewed 2026-09-06. Distinguishes predefined workflows from model-directed agents; recommends increasing complexity only where justified. Vendor engineering perspective, not an independent benchmark.

3. [OWASP Gen AI Security Project: LLM06:2025 Excessive Agency ↗](#)

Publication date not stated. Reviewed 2026-09-06. Excessive functionality, permissions and autonomy are identified as causes of excessive agency. Page belongs to the 2025 edition; exact publication day is not stated.